

Searching, Learning, and Subtopic Ordering

A Simulation-Based Analysis

Câmara, Arthur; Maxwell, David; Hauff, Claudia

DOI

[10.1007/978-3-030-99736-6_10](https://doi.org/10.1007/978-3-030-99736-6_10)

Publication date

2022

Document Version

Final published version

Published in

Advances in Information Retrieval - 44th European Conference on IR Research, ECIR 2022, Proceedings

Citation (APA)

Câmara, A., Maxwell, D., & Hauff, C. (2022). Searching, Learning, and Subtopic Ordering: A Simulation-Based Analysis. In M. Hagen, S. Verberne, C. Macdonald, C. Seifert, K. Balog, K. Nørkvåg, & V. Setty (Eds.), *Advances in Information Retrieval - 44th European Conference on IR Research, ECIR 2022, Proceedings* (pp. 142-156). (Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics); Vol. 13185 LNCS). Springer. https://doi.org/10.1007/978-3-030-99736-6_10

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.



Searching, Learning, and Subtopic Ordering: A Simulation-Based Analysis

Arthur Câmara^(✉), David Maxwell, and Claudia Hauff

Delft University of Technology, Delft, The Netherlands
{a.barbosacamara,d.m.maxwell,c.hauff}@tudelft.nl

Abstract. Complex search tasks—such as those from the *Search as Learning (SAL)* domain—often result in users developing an information need composed of several aspects. However, current models of searcher behaviour assume that individuals have an *atomic* need, regardless of the task. While these models generally work well for simpler informational needs, we argue that searcher models need to be developed further to allow for the decomposition of a complex search task into multiple aspects. As no searcher model yet exists that considers both aspects and the SAL domain, we propose, by augmenting the *Complex Searcher Model (CSM)*, the *Subtopic Aware Complex Searcher Model (SACSM)*—modelling aspects as subtopics to the user’s need. We then instantiate several agents (i.e., simulated users), with different *subtopic selection* strategies, which can be considered as different prototypical learning strategies (e.g., *should I deeply examine one subtopic at a time, or shallowly cover several subtopics?*). Finally, we report on the first large-scale simulated analysis of user behaviours in the SAL domain. Results demonstrate that the **SACSM**, under certain conditions, simulates user behaviours accurately.

1 Introduction

Over the years, a series of models¹ that describe searcher behaviour have been defined [7, 8, 23]. These often provide a post-hoc explanation of—and reasoning behind—the actions of a searcher during information seeking. One of the main drawbacks of such models is their lack of predictive capabilities: we can neither use these models to investigate what is likely to occur in different instantiations of a retrieval system; nor can we use them for simulating user behaviour.²

Indeed, models examining searcher behaviours with predictive power [2, 3, 17] have only recently been explored in the field of *Interactive Information Retrieval (IIR)*. Such models enable us to relate changing costs (e.g., the cost of examining a document) to changing searcher behaviours. Prior works employing these models have investigated how searchers interact with ranked lists [35], the impact of different browsing costs on a searcher’s behaviour [5, 22], and stopping behaviours

¹ In this paper, we refer to a *model* as a *model of user behaviour*.

² This research has been supported by *NWO* projects *SearchX* (639.022.722) and *Aspasia* (015.013.027).

on *Search Engine Results Pages (SERPs)* [31, 56]. Search topics are usually considered *atomic* in all of these prior works, with a simple information need. That is, over a *search session*³, a single topic is considered—with retrieved documents considered to be either relevant or non-relevant to that one topic. These works do not consider the different *aspects* that may constitute a wider topic. In this work, we introduce the first model of user behaviour that incorporates such thinking. More specifically, we take as a starting point the *Complex Searcher Model (CSM)* [30], a model that considers a user’s interactions throughout a search session (over multiple queries), and extend it to yield the *Subtopic-Aware Complex Searcher Model (SACSM)*—which, by considering the aspects as subtopics of a larger information need, models: (i) subtopic selection; and (ii) subtopic switching steps in the search process.

With the SACSM, we explore the effect of different subtopic switching strategies for multiple types of users within a particular domain to ground our work. We consider *Search as Learning (SAL)*, defined by Marchionini [26], as an iterative process whereby learners engage by reading, scanning and processing a large number of documents retrieved by a search system. Here, the goal is to gain knowledge about a specific learning objective. With web search engines having become an essential resource for learners [15], it is therefore vital to provide support to learners (e.g., through the form of novel interface designs [10, 45, 47] or rankers optimised for human learning [49]) that help improve their learning efficiency while searching. As a learner’s complex information needs can often be decomposed into several subtopics, a natural question to ask is *how searchers should tackle the different subtopics to learn efficiently*.

To answer this question, we present an exploratory study of the SACSM where we simulate different types of learners as *agents*⁴, and compare these to each other, examining the effect their search behaviour has on their ability to discover documents containing important keywords, as well as how they navigate throughout the subtopic space. We instantiate a series of agents that subscribe to the SACSM—with four tunable parameters that control their simulated searching behaviour: (i) **learning speed** (λ), or how fast agents incorporate novel terms into their vocabulary; (ii) **exploration** (ξ), or how willing agents are to explore each subtopic; (iii) **tolerance** (τ), or how willing an agent is to click on a search result snippet; and (iv) **subtopic switching** (φ), the strategy that agents employ to navigate through subtopics. As such, we present the first SAL study that employs simulation to examine the search behaviours of learners. By grounding a series of simulated agents with interaction data from a prior user study, we run extensive simulations of interaction to address the following research question:

RQ *How do **subtopic switching** (φ) strategies for learning-oriented search tasks affect the search behaviour of simulated agents?*

³ We consider a *search session* as interactions with a search interface, which can include the issuing of multiple queries—and the examination of multiple documents.

⁴ *Agents* are *simulated users* that are able to make judgements as to the relevancy/attractiveness of information *without* recourse to relevance information [29].

To answer **RQ**, we measure behaviours by tracking how specific measures—the *number of keywords found*, the *order of keywords found*, and *subtopic exploration*—evolve over an agent’s search sessions. We argue that to be considered effective, a strategy should allow an agent to: (i) discover as many keywords as possible in the early stages of the session; and (ii) help the agent to complete the subtopic space exploration in as few steps as possible.

The main findings of our work are: (i) subtopic switching strategies that prioritise ordering in the subtopic picking process yield improved discovery of keywords and exploration of subtopics; and (ii) the **SACSM** is enough to instantiate agents that display behaviour similar to real-world learners in a SAL context. Findings suggest that the **SACSM** is a high-quality model that provides a solid step in approximating searcher behaviours in the SAL domain. This is vital for works that rely on large-scale simulations, such as reinforcement learning for training new rankers optimised for human learning, as well as quickly evaluating new interfaces and algorithmic changes cheaply—all in a simulated environment.

2 Related Work

Models of Searcher Behaviour. Models of searcher behaviour typically fall into one of two categories: (i) descriptive models [7, 16, 20, 23, 43], allowing us to gain an intuition about the search process; and (ii) models that are expressed in more formal (mathematical) language [2, 6, 11, 17, 52, 53]. The latter category of model provides *predictive power* about why users behave in a certain way. As such, they can be used as the basis of *simulations of interaction* [4]. Here, a model of searcher behaviour that provides a credible approximation of reality can be used to ground simulations to examine what may happen under given circumstances. Despite the advantages that simulations provide, formulating such descriptive models is non-trivial. Contemporary SERPs for example are complex user interfaces, with new components (e.g., *entity cards* [37]) added all the time. In contrast, searcher models typically assume a simple SERP in the format of the traditional *ten blue links* [18]. Numerous studies have been undertaken on this more simplistic design, such as the cost of scrolling [1, 5, 42], typing [12, 40] or response time lag [28, 48].

Subtopics. The *Information Retrieval (IR)* community primarily considers the notion of subtopics from a system-centred point of view, with prior works focusing on ranking functions optimised for subtopic retrieval and result diversification [13, 21, 38, 57, 58]. Automatic subtopic (structure) extraction has also been investigated, generally based on a given starting query or document [19, 51]. The influence of subtopic characteristics on users has not been frequented in IIR. One exception is by Câmara et al. [10], who provided study participants with a list of subtopics and (visual) indicators about the extent of their subtopic exploration. The impact of subtopic ordering on users was not investigated.

SAL. We ground our work in the domain of learning which has attracted considerable attention in recent years. Beyond studies investigating how learning-oriented searches are conducted [15, 36] and how to measure learning occurring in search sessions [9, 15, 55], multiple recent studies have investigated the impact of certain user characteristics and user actions on learning during search sessions—examples include the impact of domain knowledge [39, 54], source selection strategies [25], and the cognitive abilities of users [41]. While observational studies are numerous, works proposing novel retrieval algorithms [49] and novel interface elements [10, 47] to support learning whilst searching remain sparse.

3 Subtopic-Aware Complex Searcher Model (SACSM)

For our study, we augment the CSM [32, 33] to be subtopic-aware, turning it into the SACSM. The CSM is a conceptual model of the IIR process (or a *search session*), describing the flow of activities and decisions that a searcher undertakes when interacting with a search engine. The CSM is built on the work of other conceptual models of the IIR process, such as the models of Baskaya et al. [6] and Thomas et al. [52]. Conceptual models provide us with the necessary scaffolding from which we can expand and develop the model further for a SAL context—and instantiate the model in such a way that we can run our simulations of interaction [4].

The SACSM is illustrated in Fig. 1; it includes a series of additional activities and decision points (compared to CSM) pertaining to the idea of **subtopic selection**, with novel components highlighted in blue. Key activities are represented as boxes □, with key decision points undertaken by subscribing agents represented as diamonds ◇. Upon starting at ●, a user (or, in the case of a simulation, an agent) following the SACSM will first examine the given **topic A**. SACSM then directs the agent to examine a list of the provided **subtopics B** for the given topic, before then deciding **what subtopic C** to examine in detail. From here, the agent will **consider a number of potential queries D** to issue pertaining to the selected subtopic, before **selecting a query E** to issue **F**. The agent will then obtain an ‘overview’ of the SERP **G**, and decide whether to **enter it [31] H**—and if they do, they begin to **examine a snippet I**. If the present snippet is **sufficiently attractive J**, the agent will **click the associated link K**, and **assess the document L** for usefulness and/or relevancy, before deciding to **continue on the SERP M** (and examining further snippets if so). If not, the decision to **continue with the current subtopic N** is then made. If this is the case, further queries are issued **E**—meaning that the snippet and document examination activities are repeated for the results of the new query. This also means that subtopic exploration can entail multiple queries. If the agent instead decides to **abandon the subtopic N**, they must then decide whether to **stop the search session O** altogether. This process is repeated until all subtopics have been exhausted by the agent **P**, or some other condition is met—such as running out of session time.

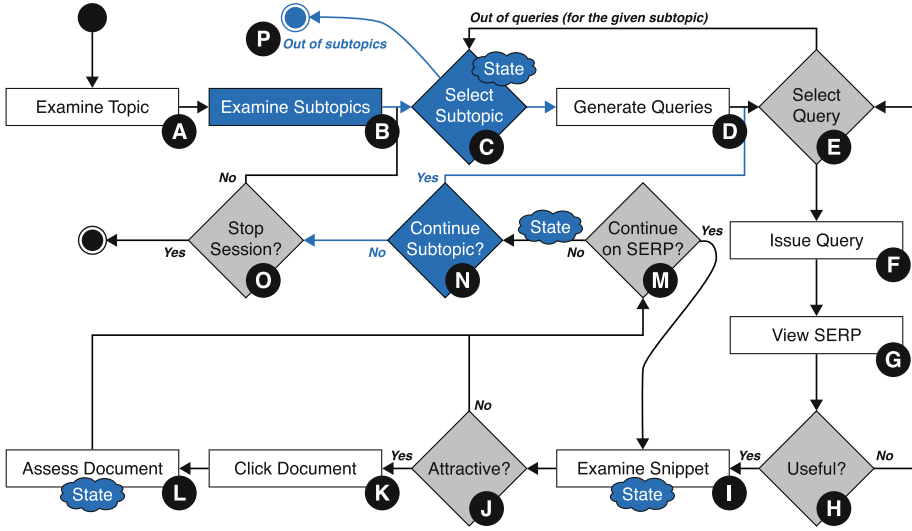


Fig. 1. The *Subtopic-Aware Complex Searcher Model (SACSM)*. Changes from the CSM are highlighted in blue. Refer to Sect. 3 for more about the sequence and shapes.

Note that compared to previous instantiations of the CSM [30–33], we have removed activities and decision points about assessing documents for relevance. Unlike simple search sessions, with *atomic* information needs, a SAL task generally has a more complex and nuanced need [15]. Therefore, we are interested in examining the *content* of retrieved documents (and thus learning from them)—not simply whether the documents themselves are considered relevant, as has been the norm for prior simulations of interaction [27].

In order to keep track of terms/concepts that are examined by agents subscribing to the SACSM (as vocabulary learning is a typical manner to measure learning gains in SAL [10,46,50]), we must also incorporate some type of *state* within it. This state model was considered in the study by Maxwell and Azzopardi [30, Fig. 3] through the *User State Model (USM)*, which “represents the user’s cognitive state”. Instead of representing the USM as a global, session-based model accumulating state and knowledge of the information examined, we consider a state model for the individual subtopics examined by agents. Each subtopic state consists of a representation of the terms observed by the agent to help them identify fundamental terms about the subtopic, which is used for query generation and determining what snippets (based on the snippet text provided by the underlying retrieval system) should be clicked on, with the corresponding document examined in more detail. Agents following this model only accrue knowledge when examining documents in full, deterministically deducing whether a document is worth examining without recourse to any relevance judgements. The state model is updated at points represented by  in Fig. 1.

4 Experimental Method

In this section, we describe the details of our instantiation of the **SACSM** and our simulations. We start by defining how we instantiate each of the components of the **SACSM** from Fig. 1, dividing them between fixed (i.e., no difference between agents) and variable components (i.e., changes between each agent). By tweaking the variable components, we can instantiate agents that simulate users with different characteristics. For example, an agent with a high λ (*how fast am I at learning new terms?*), low ξ (*how much content should I explore?*) and low τ (*how liberal am I at clicking links?*) simulates a learner that can quickly absorb new concepts, while only skimming through documents and clicking on almost all documents presented to them. We also outline our search setup, our simulation setup, as well as the datasets, and topics (and subtopics!) used.

4.1 Fixed SACSM Components

We instantiate the **SACSM** in various ways to evaluate how different **subtopic switching** strategies performed for different types of users. Although the **SACSM** has many activities and decision points to instantiate, we fixed a number of these to reduce the space that we were required to examine.

Query Generation. We use the $QS3^+$ querying strategy proposed by Maxwell and Azzopardi [30], where three query terms are selected from a language model learned from the documents the agent has already explored (plus the topic description). Previous user studies in the SAL domain [10, 47] have shown that three query terms per query is reasonable, and close to what real-world searchers use.

SERP Examination. Considered by Maxwell and Azzopardi [31], SERP examination strategies provide users with the ability to survey a SERP before committing to examining it in detail. Here, We choose to reduce the complexity of our agents (and explored space) and use the *Always Examine* approach—agents always enter the SERP and examine at least one result snippet.

User Interaction Costs. To realistically mimic how long agents should spend on each phase of their search process, we present in Table 1 the costs (in seconds) from the interaction data of a prior user study [10]. Note the high document examination cost—as participants of the user study were attempting to formulate ideas about concepts, they spent on average longer on documents when compared to other, non-SAL based studies (e.g., [31]). We also note that the total session times influence the stopping behaviours of agents since, when agents reach the time limit of their sessions, they automatically stop—regardless of the number of remaining queries to be issued, as generated by the $QS3^+$ strategy.

Snippet-Level Stopping Strategies. Different *snippet-level stopping strategies* can be employed, generally classified between *fixed* (i.e., the agent will evaluate snippets until a certain depth) or *adaptive* strategies (i.e., the number of snippets evaluated may change depending on factors like agent state, presented snippet

Table 1. Interaction costs grounding our agents, as derived from Câmara et al. [10].

Time required to...	Value (in seconds)
...issue a query	9.42
...examine a SERP	2.00
...examine a result snippet	3.00
...examine a document	80.00
Total session time	2400

content, etc.). For our agents, we use only a fixed snippet-level stopping strategy, where agents examine snippets to a depth of 10. This is a reasonable depth to examine to, and avoids issues with SERP pagination.

4.2 Variable SACSM Components

For this study, agents can be instantiated using four variables, according to the type of user to be simulated. An overview of these variables is presented in Fig. 2.

Subtopic Switching (φ). We propose four different strategies for agents to select and switch between subtopics during their search sessions. These implement the **Select Subtopic** decision point, as shown in Fig. 1. To determine whether agents have explored a subtopic sufficiently, we use a method similar to the approach outlined by Câmara et al. [10] for tracking subtopic exploration. Each clicked document is embedded using **SBERT** [44] and compared—using the dot product—to pre-computed embeddings for each subtopic of the current topic, as extracted from their *Wikipedia* articles⁵. Therefore, each document clicked by an agent will update an internal state tracker for each subtopic, summing how much the agent ‘explored’ each subtopic. We evaluate four strategies.

- **Greedy** For this strategy, an agent examines each subtopic in turn, according to the order provided by the respective *Wikipedia* article, only deciding to move to the next subtopic when they have achieved a certain level of progress. Intuitively, this would be the most rational type of user, since they follow a subtopic ordering that is optimised for human understanding (i.e., the order comes from a *Wikipedia* page). In other words, they will attempt to master one subtopic before moving to the next (prescribed) topic.
- **Greedy-Skip** Instead of the above, an agent subscribing to **Greedy-Skip** moves to the next subtopic with the *next lowest completion value*. This instantiated agent attempts to minimise the number of documents to be read by querying in a domain with lesser knowledge.
- **Reverse** This strategy is similar to **Greedy**, but the agent examines the subtopics in reverse order as presented from the corresponding *Wikipedia* article. The rationale here is that an agent attempts to game the system by first learning the most complex subtopics *before* moving to easier ones.

⁵ Refer to Sect. 4.3 for more information on the use of *Wikipedia* articles.

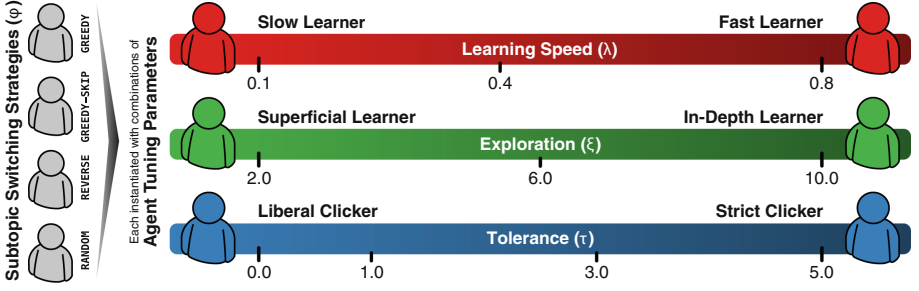


Fig. 2. Overview of the four variable parameters for instantiating simulated agents.

- **Random** This strategy randomly selects a new subtopic after each query, with no predefined order. This strategy models a non-rational learner, and serves as a lower bound for our experiments.

Learning Speed (λ). This parameter is the same λ from the language model proposed by Maxwell and Azzopardi [30]. It controls how much an agent relies on their acquired knowledge (i.e., novel terms) when considering if a given snippet should be clicked or not. The language model is updated every time the agent clicks on a relevant snippet. In addition, a *Maximum Likelihood Estimator* [34] is used for deciding if a given snippet is attractive or not. An agent with a low λ gives lower weights to terms learned during the session, simulating a *slow learner*. An agent with a higher λ , in turn, mimics a user that quickly incorporate new terms, being a *fast learner*. In our simulations, we use $\lambda \in [0.1, 0.4, 0.8]$.

Exploration (ξ). This parameter controls how much the agent should explore a subtopic before being satisfied by what it has ‘*learnt*’. A lower number implies that such an agent that is only ‘skimming’ through the topic, inspecting only a few documents per subtopic. In contrast, a higher value implies that such an agent is willing to explore deeper within each subtopic. For the simulations reported in this paper, we trial $\xi \in [2.0, 6.0, 10.0]$.

Tolerance (τ). Finally, this parameter is the threshold that controls how attractive a snippet should be to be clicked [34]. An agent with low τ is a *strict clicker*, clicking on fewer, ‘safer’ snippets, while a higher τ implies a *liberal clicker* agent, more willing to explore. In our simulations, we trial $\tau \in [0.0, 1.0, 3.0, 5.0]$.

4.3 Simulation Setup

The setup of our experiments follows that of the user study presented by Câmara et al. [10]. Here, we use the same eight topics extracted from the TREC CAR 2017 dataset [14], as shown in Table 2. Subtopics were also derived from the TREC CAR dataset: they were extracted from first-level headings of the respective *Wikipedia* articles as they were in December 2017—the dataset’s creation.

Table 2. # of subtopics and distinct keywords (KW) for each topic. For each subtopic, we determine the ten KWs with higher TF-IDF on the respective subtopic section on Wikipedia. A KW may appear in the top ranks of several subtopics. KW difficulty is given by the age-of-acquisition, as proposed by Kuperman et al. [24].

Topic	#Subtopics	#Unique KWs	KW Difficulty
Ethics	6	49	10.85
Genetically Modified Organism	5	33	9.97
Noise-Induced Hearing Loss	8	56	8.85
Subprime Mortgage Crisis	8	52	9.81
Radiocarbon Dating	4	35	9.77
Business Cycle	4	32	10.70
Irritable Bowel Syndrome	10	72	9.88
Theory of Mind	8	67	9.63

Our study uses the *Bing Search API* to provide a ranking for queries issued by real-world users and our simulated agents. We used a manually curated blocklist⁶ of URLs serving *Wiki*-style clones to filter results returned from the Bing API to prevent agents from encountering a single page that would give them all the information on all subtopics at once. This encourages agents to examine multiple documents and issue multiple queries to find information pertinent to their learning task. To match with our stopping strategy (see Sect. 4.1), 10 results per page were presented to agents.

5 Results

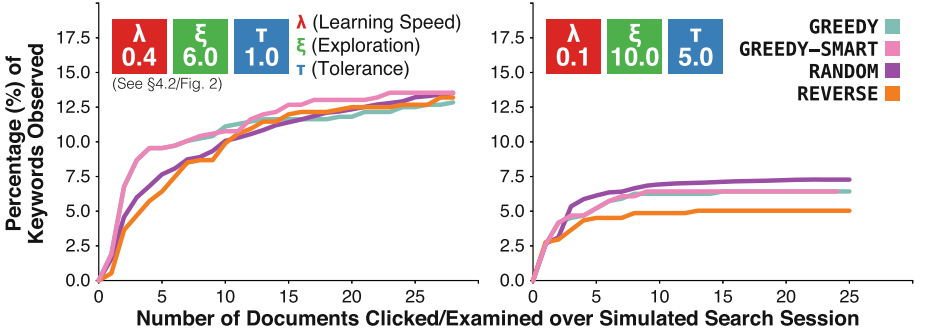
By combining all values of ξ , λ , τ and φ , we instantiate 144 unique agents (using a modified version *SimIIR* [30]), and run each agent over all of the topics shown in Table 2. Our version of *SimIIR*—together with the raw outputs of our simulations—are available at <https://github.com/ArthurCamara/simiir-subtopics/>. With some methods being non-deterministic, each agent was run ten times—with the average reported. In total, we ran a total of 11,520 simulations. We show representative examples for each set of measures. In all plots, the x axes denote how many documents the agent examined during a search session. While values on y axes may seem low, they are averaged over a large number of simulations with varying degrees of complexity.

Table 3 shows the average value for key measures over all agents of each φ , over the eight topics. In the first row, we also show the measures from the *FEEDBACK_{SC}* cohort from Câmara et al. [10]. Our simulated agents are similar on these measures compared to how real-world learners would behave, with a similar number of queries, snippets examined, and documents clicked. While a high deviation is expected, recall that this is an average of 288 agents with a large

⁶ <https://github.com/ArthurCamara/CHIIR21-SAL-Scaffolding/blob/master/data/blocklist.txt> (All URLs last accessed January 18th, 2022.).

Table 3. Overview of (average) measures across agents and subtopic switching strategies, and real learners extracted from the $\text{FEEDBACK}_{\text{SC}}$ cohort from Câmara et al. [10].

Strategy (φ)	#Queries Issued	#Snippets Examined	#Documents Clicked
$\text{FEEDBACK}_{\text{SC}}$ (N=36)	11.86(± 7.60)	152.44(± 84.23)	18.50(± 9.56)
Greedy	13.05(± 14.93)	133.37(± 176.29)	21.32(± 7.60)
Greedy-Skip	13.05(± 15.13)	133.23(± 177.26)	21.44(± 8.88)
Reverse	12.01(± 14.55)	123.28(± 173.34)	21.42(± 8.78)
Random	13.03(± 16.07)	117.61(± 155.80)	21.82(± 8.30)

**Fig. 3.** Accumulated percentage of the keywords seen for two different agents (averaged over all topics, weighted by number of keywords) with varying ξ , λ and τ .

variation on their parameters (compared to only 36 real-world learners). Results show that our agents are indeed similar to real-world learners. To address **RQ**, we break down our analysis further into three sub-questions.

How Many Keywords Can the Agents Find? To measure how well the agents can find documents with a high concentration of potentially valuable keywords⁷, we extracted ten keywords for each subtopic from their respective paragraphs from the topic’s Wikipedia article.⁸ To do this, we begin by ranking all terms from their portions of the articles (excluding stopwords) by their TF-IDF, with the IDF computed over the whole TREC CAR Wikipedia dump—and selecting the top 10 terms as keywords. We use this subtopic-wise approach (instead of extracting keywords from the whole article) to ensure a fair distribution of keywords over all subtopics, providing a less biased overview of how the agent is performing over the topic. Therefore, each topic has a different number of keywords, reflected by the number of subtopics it contains. Table 2 shows how many subtopics and unique keywords each topic contains. This setup is similar to previous SAL user studies [10, 46, 47], where study participants were asked to

⁷ As noted in Sect. 3, we do not have explicit relevance judgements.

⁸ As an example, the following are extracted keywords for the topic *Ethics*: **ethical**, **ontology**, **propositions**, **consequentialism**, **normative** and **principles**.

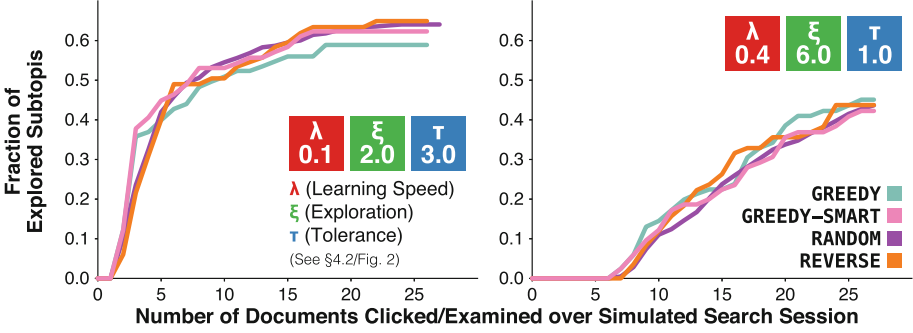


Fig. 4. Fraction of fully explored (i.e. agent reached ξ value) for two different agents (averaged over all topics, weighted by number of keywords) with varying ξ , λ and τ .

define a list of concepts before and after their search session to evaluate their knowledge gain. We can mimic this setup throughout the entirety of an agent’s search session by requiring the keyword to appear at least a few times (in our case, five) in the documents ‘read’ by the agent.

For two agents, Fig. 3 shows how many keywords each approach for φ discovers during their search sessions after reading a certain number of documents. At the beginning of the search sessions, we observed that agents instantiated with **Greedy** and **Greedy-Skip** strategies found keywords faster than agents with **Random** or **Reverse**. However, this difference diminishes over time. This is expected, since the subtopics ordering comes from Wikipedia articles, which are optimised for human understanding. Therefore, an agent that searches for subtopics in order has a higher probability of encountering documents with a higher number of keywords earlier in the session when compared with one that does not. We can also note that **Random** with higher τ found more unique keywords in total, given their high probability of clicking in any document.

Are the Agents Exploring Enough of the Subtopics? Another way to measure how the agents behave is by investigating how their internal *Subtopic Trackers* evolve during the session (as explained in Sect. 3). If an agent can reach ξ for a given subtopic in a few documents, we can infer that they could quickly find documents related to that subtopic. Figure 4 shows a similar trend to that observed previously, with agents using **Greedy** and **Greedy-Skip** strategies clicking on documents that advance their internal tracking faster. This implies that these strategies effectively lead the agents towards better documents faster.

Are the Agents Following the Order of the Subtopics? While the previous measures show that the agents are indeed effective in finding documents related to the topic, they fail to incorporate another essential learning feature, namely that keywords have dependencies between them. We assume that, for an agent to comprehend what a keyword means entirely, they have to comprehend at least some other, more basic concepts related to the topic at hand. Therefore, an agent that can find documents so that they will encounter more primary

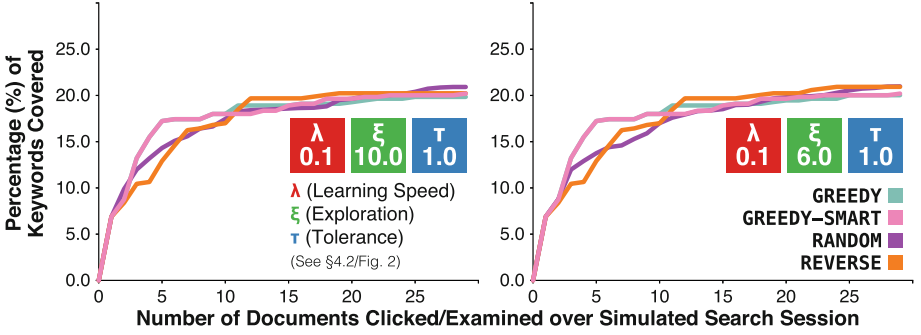


Fig. 5. Fraction of keywords properly ‘learned’ by two different agents (averaged over all topics, weighted by number of keywords) with varying ξ , λ and τ .

keywords for the topic (i.e., that appear earlier in the Wikipedia article) earlier in the session before facing more complex keywords (i.e., that appear later in the Wikipedia article) is more desirable for a SAL environment. As an example, consider the keyword *consequentialist* for the topic *Ethics*. Before an agent can adequately understand what it means in this context, they probably need to understand other concepts, like *virtue* and *morality*. Therefore, for this analysis, we consider a keyword to be ‘learned’ after the agent has already encountered a certain number of keywords that appear prior to it in the original Wikipedia article. To account for possible noises in our keyword extraction method, we define this number as 50% of the keywords seen prior to the current one (e.g., the keyword *consequentialist* is the 19th out of 49 keywords to appear in the list of extracted terms for the topic *Ethics*). Therefore, we only consider a given keyword as learned after the agent has learned at least 30% of the prior terms.⁹ As seen in Fig. 5, we see similar behaviour to the one observed above, with **Greedy** and **Greedy-Skip** outperforming **Random** and **Reverse**, with the difference slowly disappearing throughout the search session. Again, almost all agents repeat this behaviour. These results show that our simulations are close to real users and that there is a clear difference between strategies, with **Greedy** and **Greedy-Skip** following the logical structure of the subtopics, and generally being better strategies for agents exploring the subtopic space. Consequently, these should be taken into account when simulating agents for SAL scenarios.

6 Conclusions

We have proposed a novel user model for simulating agents focused on SAL tasks: the *Subtopic-Aware Complex Searcher Model*, **SACSM**. Recalling our original research question which considered how different subtopic switching strategies (φ) affected the behaviour of simulated agents, we show that strategies that

⁹ This number was decided experimentally, as it showed to be the best to distinguish between the different methods trialled.

mimic a rational user (i.e., **Greedy** and **Greedy-Skip**) are more effective at *finding keywords*, *exploring subtopics* and *following subtopic structure* when compared to other strategies. With 1,520 simulations, our study is the first (to the best of our knowledge) that focuses on simulated agents for Search as Learning, enabling future works in both SAL and IIR that may require large quantities of user data, such as Reinforcement Learning models and studies on how changes in the search system may impact the behaviour of learners. To further help research efforts, we also make public our implementation of the **SACSM**, built on top of the already established **SimIIR** framework.

References

1. Albers, M., Kim, L.: Information design for the small-screen interface: an overview of web design issues for personal digital assistants. *Tech. Comm.* **49**(1), 45–60 (2002)
2. Azzopardi, L.: The economics in interactive information retrieval. In: *Proceedings of 34th ACM SIGIR*, pp. 15–24 (2011)
3. Azzopardi, L.: Modelling interaction with economic models of search. In: *Proceedings of 37th ACM SIGIR*, pp. 3–12 (2014)
4. Azzopardi, L., Järvelin, K., Kamps, J., Smucker, M.D.: Report on the SIGIR 2010 workshop on the simulation of interaction. *SIGIR Forum* **44**(2), 35–47 (2011)
5. Azzopardi, L., Zucco, G.: Two scrolls or one click: a cost model for browsing search results. In: *Proceedings of 38th ECIR*, pp. 696–702 (2016)
6. Baskaya, F., Keskustalo, H., Järvelin, K.: Modeling behavioral factors in interactive information retrieval. In: *Proceedings of 22nd ACM CIKM*, pp. 2297–2302 (2013)
7. Bates, M.J.: The design of browsing and berrypicking techniques for the online search interface. *Online Review* (1989)
8. Belkin, N.J.: The cognitive viewpoint in information science. *J. Info. Sci.* **16**(1), 11–15 (1990)
9. Bhattacharya, N., Gwizdka, J.: Relating eye-tracking measures with changes in knowledge on search tasks. In: *Proceedings of 13th ACM ETRA*, pp. 1–5 (2018)
10. Câmara, A., Roy, N., Maxwell, D., Hauff, C.: Searching to learn with instructional scaffolding. In: *Proceedings of 6th ACM CHIIR*, pp. 209–218 (2021)
11. Carterette, B., Kanoulas, E., Yilmaz, E.: Simulating simple user behavior for system effectiveness evaluation. In: *Proceedings of the 20th ACM CIKM*, pp. 611–620 (2011)
12. Crescenzi, A., Kelly, D., Azzopardi, L.: Time pressure and system delays in information search. In: *Proceedings of 38th ACM SIGIR*, pp. 767–770 (2015)
13. Dai, W., Srihari, R.: Minimal document set retrieval. In: *Proceedings of 14th ACM CIKM*, pp. 752–759 (2005)
14. Dietz, L., Verma, M., Radlinski, F., Craswell, N.: TREC complex answer retrieval overview. *TREC* (2017)
15. Eickhoff, C., Teevan, J., White, R., Dumais, S.: Lessons from the journey: a query log analysis of within-session learning. In: *Proceedings of 7th ACM WSDM*, pp. 223–232 (2014)
16. Ellis, D.: Modeling the information-seeking patterns of academic researchers: a grounded theory approach. *Libr. Q.* **63**(4), 469–486 (1993)
17. Fuhr, N.: A probability ranking principle for interactive information retrieval. *Inf. Retrieval* **11**(3), 251–265 (2008)

18. Hearst, M.: Search User Interfaces. Cambridge University Press (2009)
19. Hearst, M., Plaunt, C.: Subtopic structuring for full-length document access. In: Proceedings of 16th ACM SIGIR, pp. 59–68 (1993)
20. Ingwersen, P., Järvelin, K.: The Turn: Integration of Information Seeking and Retrieval in Context, vol. 18, p. 448. Springer, Dordrecht (2006). <https://doi.org/10.1007/1-4020-3851-8>
21. Jiang, Z., Wen, J.R., Dou, Z., Zhao, W.X., Nie, J.Y., Yue, M.: Learning to diversify search results via subtopic attention. In: Proceedings of 40th ACM SIGIR, pp. 545–554 (2017)
22. Kashyap, A., Hristidis, V., Petropoulos, M.: FACeTOR: cost-driven exploration of faceted query results. In: Proceedings of 19th ACM CIKM, pp. 719–728 (2010)
23. Kuhlthau, C.C.: Developing a model of the library search process: cognitive and affective aspects, pp. 232–242 (1988)
24. Kuperman, V., Stadthagen-Gonzalez, H., Brysbaert, M.: Age-of-acquisition ratings for 30,000 English words. *Behav. Res. Methods* **44**(4), 978–990 (2012)
25. Liu, C., Song, X.: How do information source selection strategies influence users' learning outcomes. In: Proceedings of 3th ACM CHIIR, pp. 257–260 (2018)
26. Marchionini, G.: Exploratory search: from finding to understanding. *Commun. ACM* **49**(4), 41–46 (2006)
27. Maxwell, D.: Modelling Search and Stopping in Interactive Information Retrieval. Ph.D. thesis, University of Glasgow, Scotland (2019)
28. Maxwell, D., Azzopardi, L.: Stuck in traffic: how temporal delays affect search behaviour. In: Proceedings of 5th IIX, pp. 155–164 (2014)
29. Maxwell, D., Azzopardi, L.: Agents, simulated users and humans: an analysis of performance and behaviour. In: CIKM, pp. 731–740. ACM (2016)
30. Maxwell, D., Azzopardi, L.: Simulating interactive information retrieval: SimIIR: a framework for the simulation of interaction. In: Proceedings of 39th ACM SIGIR, pp. 1141–1144 (2016)
31. Maxwell, D., Azzopardi, L.: Information scent, searching and stopping. In: Proceedings of 40th ECIR, pp. 210–222 (2018)
32. Maxwell, D., Azzopardi, L., Järvelin, K., Keskustalo, H.: An initial investigation into fixed and adaptive stopping strategies. In: Proceedings of 38th ACM SIGIR, pp. 903–906 (2015)
33. Maxwell, D., Azzopardi, L., Järvelin, K., Keskustalo, H.: Searching and stopping: an analysis of stopping rules and strategies. In: Proceedings of 24th ACM CIKM, pp. 313–322 (2015)
34. Meij, E., Weerkamp, W., de Rijke, M.: A query model based on normalized log-likelihood. In: Proceedings of 18th ACM CIKM, pp. 1903–1906 (2009)
35. Moffat, A., Scholer, F., Thomas, P.: Models and metrics: IR evaluation as a user process. In: Proceedings of 17th ADCS, pp. 47–54 (2012)
36. Moraes, F., Putra, S.R., Hauff, C.: Contrasting search as a learning activity with instructor-designed learning. In: Proceedings of 27th ACM CIKM, pp. 167–176 (2018)
37. Navalpakkam, V., Jentzsch, L., Sayres, R., Ravi, S., Ahmed, A., Smola, A.: Measurement and modeling of eye-mouse behavior in the presence of nonlinear page layouts. In: Proceedings of the 22nd WWW, pp. 953–964 (2013)
38. Nguyen, T.N., Kanhabua, N.: Leveraging dynamic query subtopics for time-aware search result diversification. In: Proceedings of 36th ECIR, pp. 222–234 (2014)
39. O'Brien, H.L., Kampen, A., Cole, A.W., Brennan, K.: The role of domain knowledge in search as learning. In: Proceedings of 5th ACM CHIIR, pp. 313–317 (2020)

40. Ong, K., Järvelin, K., Sanderson, M., Scholer, F.: Qwerty: the effects of typing on web search behavior. In: Proceedings of 3rd ACM CHIIR, pp. 281–284 (2018)
41. Pardi, G., von Hoyer, J., Holtz, P., Kammerer, Y.: The role of cognitive abilities and time spent on texts and videos in a multimodal searching as learning task. In: Proceedings of 5th ACM CHIIR, pp. 378–382 (2020)
42. Peytchev, A., Couper, M.P., McCabe, S.E., Crawford, S.D.: Web survey design: paging versus scrolling. *Intl. J. Pub. Opin. Q.* **70**(4), 596–607 (2006)
43. Pirolli, P., Card, S.: Information foraging. *Psychol. Rev.* **106**(4), 643 (1999)
44. Reimers, N., Gurevych, I.: Sentence-BERT: sentence embeddings using siamese BERT-networks. arXiv preprint [arXiv:1908.10084](https://arxiv.org/abs/1908.10084) (2019)
45. Roy, N., Câmara, A., Maxwell, D., Hauff, C.: Incorporating widget positioning in interaction models of search behaviour. In: Proceedings of 11th ACM ICTIR, pp. 53–62 (2021)
46. Roy, N., Moraes, F., Hauff, C.: Exploring users’ learning gains within search sessions. In: Proceedings of 5th ACM CHIIR, pp. 432–436 (2020)
47. Roy, N., Torre, M.V., Gadiraju, U., Maxwell, D., Hauff, C.: Note the highlight: incorporating active reading tools in a search as learning environment. In: Proceedings of 6th ACM CHIIR, pp. 229–238 (2021)
48. Schurman, E., Brutlag, J.: Performance related changes and their user impact. In: O’Reilly Velocity Conference (2009)
49. Syed, R., Collins-Thompson, K.: Retrieval algorithms optimized for human learning. In: Proceedings of 40th ACM SIGIR, pp. 555–564 (2017)
50. Syed, R., Collins-Thompson, K.: Exploring document retrieval features associated with improved short-and long-term vocabulary learning outcomes. In: Proceedings of 3th ACM CHIIR, pp. 191–200 (2018)
51. Takaki, T., Fujii, A., Ishikawa, T.: Associative document retrieval by query subtopic analysis and its application to invalidity patent search. In: Proceedings of 13th ACM CIKM, pp. 399–405 (2004)
52. Thomas, P., Moffat, A., Bailey, P., Scholer, F.: Modeling decision points in user search behavior. In: Proceedings of 5th IiX, pp. 239–242 (2014)
53. Wang, J., Zhu, J.: Portfolio theory of information retrieval. In: Proceedings of 32nd ACM SIGIR, pp. 115–122 (2009)
54. Wildemuth, B.M.: The effects of domain knowledge on search tactic formulation. *J. Am. Soc. Info. Sci. Tech.* **55**(3), 246–258 (2004)
55. Wilson, M.J., Wilson, M.L.: A comparison of techniques for measuring sensemaking and learning within participant-generated summaries. *J. Am. Soc. Info. Sci. Tech.* **64**(2), 291–306 (2013)
56. Wu, W.C., Kelly, D., Sud, A.: Using information scent and need for cognition to understand online search behavior. In: Proceedings of 37th ACM SIGIR, pp. 557–566 (2014)
57. Zhai, C., Cohen, W.W., Lafferty, J.: Beyond independent relevance: methods and evaluation metrics for subtopic retrieval. *SIGIR Forum* **49**(1), 2–9 (2015)
58. Zuccon, G., Azzopardi, L., Hauff, C., van Rijsbergen, C.K.: Estimating interference in the QPRP for subtopic retrieval. In: Proceedings of 33rd ACM SIGIR, pp. 741–742 (2010)